

PERBANDINGAN KINERJA METODE NAIVE BAYES DAN SUPPORT VECTOR MACHINE (SVM) DALAM ANALISIS SENTIMEN PUBLIK TERHADAP PROGRAM MAKAN BERGIZI GRATIS (MBG)

Oleh :

Mohamad Yachob Oktavian (221080200066)

Dosen Pembimbing :

Yulian Findawati, ST., M.MT.

Dosen Penguji

Suhendro Busono, S.ST., M.Kom.

Uce Indahyanti S.Kom.M.Kom., Dr.

Program Studi Informatika

Fakultas Sains dan Teknologi

Universitas Muhammadiyah Sidoarjo

2026

Latar Belakang

Program Makan Bergizi Gratis (MBG) merupakan program pemerintah untuk mengatasi stunting dan malnutrisi yang memunculkan berbagai tanggapan masyarakat di media sosial Twitter (X). Banyaknya tweet membuat analisis manual menjadi tidak efektif, sehingga diperlukan metode machine learning. Oleh karena itu, penelitian ini membandingkan kinerja Naive Bayes dan Support Vector Machine (SVM) menggunakan ekstraksi fitur TF-IDF dalam mengklasifikasikan sentimen publik terhadap Program MBG.

Penelitian Terdahulu

Berdasarkan beberapa penelitian sebelumnya, metode Support Vector Machine (SVM) umumnya memperoleh akurasi yang lebih tinggi dibandingkan Naive Bayes pada berbagai topik analisis sentimen, seperti COVID-19, RUU Kesehatan, BPJS, Mobil Listrik, SatuSehat, dan IKN. Hasil tersebut menjadi dasar untuk menguji kembali apakah SVM juga memberikan performa terbaik pada kasus Program Makan Bergizi Gratis.

Peneliti	Domain	SVM	NB	Selisih
Sujadi et al. (2022)	COVID-19	81,6%	78,3%	3,3%
Widyanto et al. (2023)	RUU Kesehatan	87%	82%	5%
Matheos (2024)	SatuSehat App	87,95%	81,65%	6,3%
Amandasari (2025)	Layanan BPJS	79,32%	61,63%	17,69%
Ningsih et al. (2024)	Mobil Listrik	70,82%	63,02%	7,8%
Supian et al. (2024)	IKN	94%	91%	3%

Analisis Gap

- Belum ada penelitian spesifik tentang sentimen program pemerintah berskala besar yang melibatkan berbagai stakeholder.
- Penelitian sebelumnya fokus pada aplikasi kesehatan digital dan pelayanan kesehatan publik.
- Evaluasi kinerja cenderung hanya fokus pada akurasi tanpa eksplorasi mendalam precision, recall, dan F1-score.
- Penelitian MBG sebelumnya menggunakan data fase awal dan pendekatan rule-based tanpa machine learning.

Rumusan Masalah

1. Bagaimana penerapan metode Naive Bayes dan Support Vector Machine (SVM) dalam mengklasifikasikan sentimen publik (positif, negatif, netral) terhadap Program MBG berdasarkan data Twitter?
2. Bagaimana perbandingan kinerja kedua metode ditinjau dari metrik akurasi, presisi, recall, dan F1-score?

Tujuan Penelitian

Penelitian ini bertujuan mengimplementasikan metode Naive Bayes dan Support Vector Machine (SVM) untuk mengklasifikasikan sentimen publik terhadap Program MBG serta membandingkan performa kedua metode berdasarkan akurasi, presisi, recall, dan F1-Score sehingga dapat diketahui metode yang paling efektif digunakan.

Batasan Masalah

- Data tweet berbahasa Indonesia dari platform Twitter (X) tentang program MBG periode Januari-Desember 2025.
- Klasifikasi sentimen dibatasi pada tiga kelas: positif, negatif, dan netral.
- Metode yang dibandingkan: Naive Bayes dan SVM dengan teknik feature extraction TF-IDF.
- Evaluasi menggunakan metrik akurasi, presisi, recall, dan F1-score.
- Implementasi menggunakan Python 3.8+ dengan library Scikit-learn, Sastrawi, dan Pandas/NumPy.

Metodologi Penelitian

Penelitian diawali dengan pengumpulan data menggunakan web scraping, kemudian dilakukan proses data cleaning dan preprocessing yang meliputi case folding, cleaning, tokenisasi, stopword removal, serta stemming menggunakan Sastrawi. Selanjutnya data diubah menjadi fitur TF-IDF, dibagi menjadi data latih dan data uji dengan rasio 80:20, kemudian dilatih menggunakan metode Naive Bayes dan SVM sebelum dilakukan evaluasi.

Data Penelitian

Dataset diperoleh dari Twitter (X) menggunakan Selenium dengan periode pengambilan data Januari hingga Desember 2025. Sebanyak **1.109 tweet** berhasil dikumpulkan dan setelah proses pembersihan diperoleh **1.050 tweet valid**. Distribusi sentimen relatif seimbang, yaitu Netral **34,2%**, Positif **33,0%**, dan Negatif **32,8%**, sehingga cukup baik digunakan dalam proses klasifikasi.

Preprocessing Data

- 1. Case Folding: Mengubah semua teks menjadi huruf kecil untuk konsistensi.
- 2. Cleaning: Menghapus URL, mention (@), hashtag, karakter khusus, dan angka tidak relevan.
- 3. Tokenisasi: Memecah teks menjadi token/kata-kata individual.
- 4. Stopword Removal: Menghilangkan kata-kata umum yang tidak bermakna (dan, yang, di, dll.).
- 5. Stemming: Mengubah kata ke bentuk dasar menggunakan library Sastrawi.
- 6. TF-IDF: Mengubah teks menjadi representasi numerik berdasarkan frekuensi dan bobot kata.

Naive Bayes

- Algoritma klasifikasi berbasis probabilitas yang bekerja dengan menerapkan teorema Bayes dengan asumsi independensi antar fitur.
- Formula Teorema Bayes: $P(C|X) = [P(X|C) \times P(C)] / P(X)$
- Kelebihan: Sederhana dan cepat, efektif untuk dataset besar, tidak memerlukan data latih yang sangat besar.
- Kekurangan: Asumsi independensi tidak selalu terpenuhi, sensitif terhadap zero frequency, kurang efektif menangkap hubungan kompleks antar kata.

Support Vector Machine (SVM)

- Algoritma yang mencari hyperplane optimal untuk memisahkan data ke dalam kelas berbeda dengan margin maksimum. Menggunakan kernel RBF untuk data non-linear.
- Fungsi Keputusan: $f(x) = \text{sign}(w \cdot x + b)$
- Kelebihan: Efektif pada data berdimensi tinggi, robust terhadap overfitting, mampu menangani data tidak seimbang, generalisasi baik.
- Kekurangan: Waktu pelatihan lebih lama, memerlukan pemilihan kernel dan parameter yang tepat, sulit diinterpretasi.

Metrik Evaluasi

- Akurasi: $(TP + TN) / (TP + TN + FP + FN)$ - Proporsi prediksi yang benar dari keseluruhan data.
- Presisi: $TP / (TP + FP)$ - Proporsi prediksi positif yang benar.
- Recall: $TP / (TP + FN)$ - Proporsi data positif yang berhasil teridentifikasi oleh model.
- F1-Score: $2 \times (Precision \times Recall) / (Precision + Recall)$ - Harmonic mean antara presisi dan recall.

Hasil Ekstraksi Fitur & Pelatihan Model

Proses ekstraksi menggunakan TF-IDF menghasilkan 2.920 fitur unik dari seluruh tweet. Dataset kemudian dibagi menjadi 840 data latih dan 210 data uji. Hasil pelatihan menunjukkan bahwa Naive Bayes memiliki waktu training dan testing yang jauh lebih cepat dibandingkan SVM, yaitu sekitar 52 kali lebih cepat pada proses training.

Perbandingan Metrik Evaluasi (NB vs SVM)

Metrik	Naive Bayes	SVM	Selisih (NB-SVM)
Akurasi	70,48%	68,57%	+1,91%
Presisi	71,14%	70,10%	+1,04%
Recall	70,48%	68,57%	+1,91%
F1-Score	69,73%	68,87%	+0,86%
Waktu Training	0,0047 dtk	0,2447 dtk	NB 52× lebih cepat
Waktu Testing	0,0006 dtk	0,0371 dtk	NB 62× lebih cepat

Hasil evaluasi menunjukkan bahwa Naive Bayes memperoleh nilai sedikit lebih tinggi dibandingkan SVM pada seluruh metrik yang diuji, yaitu akurasi 70,48%, presisi 71,14%, recall 70,48%, dan F1-Score 69,73%. Meskipun selisih akurasinya hanya 1,91%, Naive Bayes memiliki keunggulan yang lebih jelas dari sisi efisiensi komputasi.

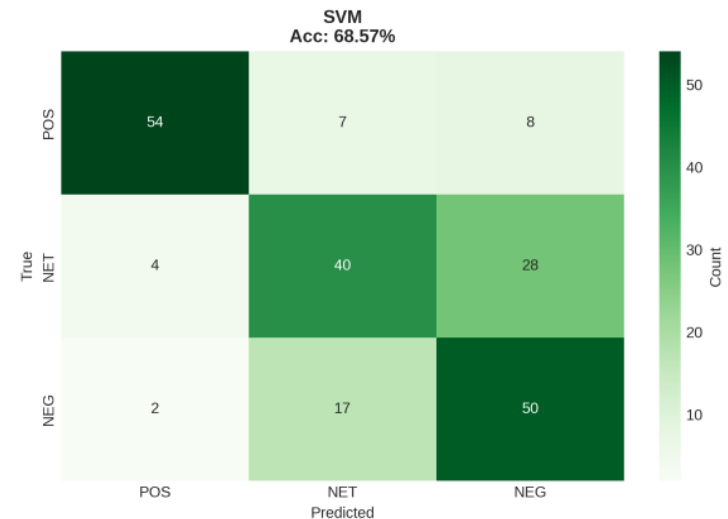
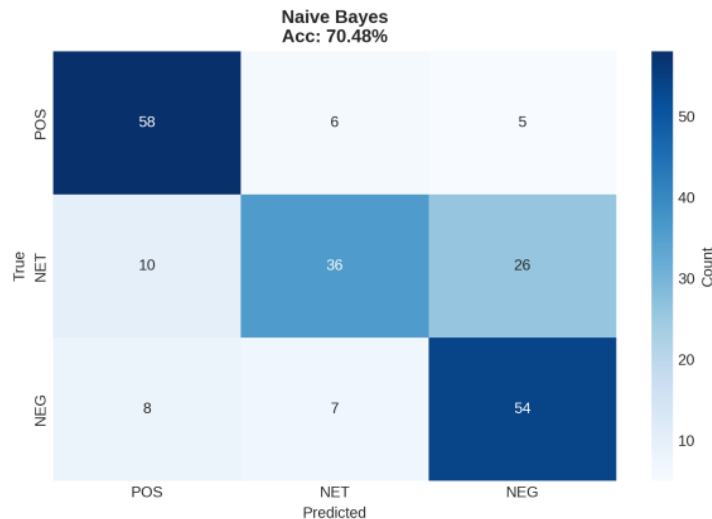
Analisis per Kelas Sentimen (Tabel 4)

Kelas	NB Precision	NB Recall	NB F1	SVM Precision	SVM Recall	SVM F1
Positif	76,32%	84,06%	80,00%	90,00%	78,26%	83,72%
Netral	73,47%	50,00%	59,50%	62,50%	55,56%	58,82%
Negatif	63,53%	78,26%	70,13%	58,14%	72,46%	64,52%

Berdasarkan hasil evaluasi setiap kelas, SVM memiliki nilai presisi tertinggi pada kelas positif, sedangkan Naive Bayes memperoleh recall yang lebih baik pada kelas positif serta performa yang lebih baik pada kelas negatif. Sementara itu, kelas netral menjadi kategori yang paling sulit diklasifikasikan oleh kedua metode karena karakteristik bahasanya yang tidak memiliki polaritas emosi yang kuat.

Confusion Matrix & Pola Kesalahan

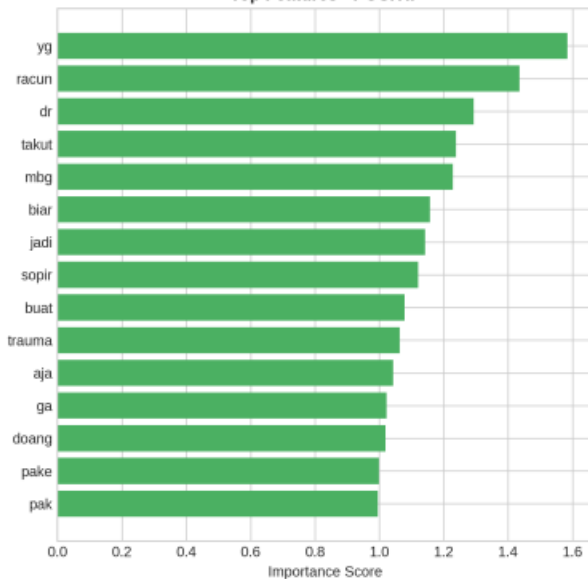
- Kesalahan paling menonjol pada kedua model: misklasifikasi kelas Netral menjadi Negatif (NB: 26 dari 72 tweet Netral; SVM: 28 kesalahan serupa).
- Hal ini mengindikasikan tumpang tindih fitur linguistik antara kelas Netral dan Negatif, karena tweet netral sering membahas isu negatif (mis. kasus keracunan) secara deskriptif tanpa ekspresi emosi eksplisit.
- Kedua model tampil terbaik pada kelas Positif: NB mengklasifikasikan 58 dari 69 tweet benar (84,06%), SVM 54 dari 69 (78,26%).



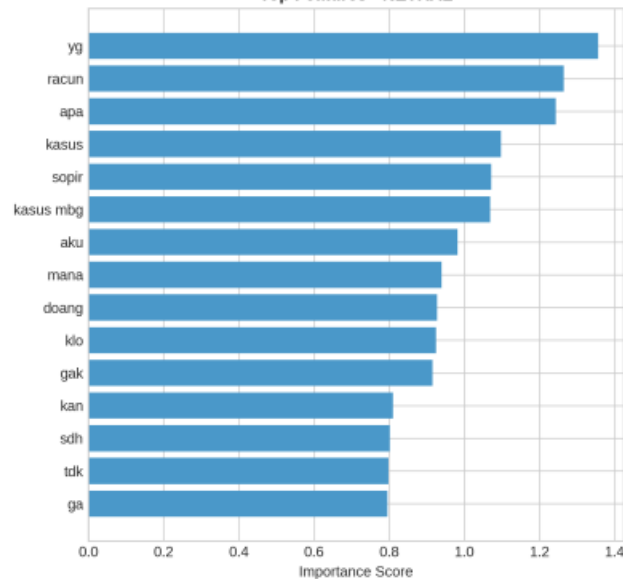
Pembahasan

Distribusi sentimen yang relatif seimbang menunjukkan bahwa masyarakat masih berada pada tahap evaluasi terhadap Program MBG. Berbeda dengan sebagian besar penelitian terdahulu, penelitian ini menunjukkan bahwa Naive Bayes sedikit lebih unggul dibandingkan SVM, yang diduga dipengaruhi oleh distribusi kelas yang seimbang, ukuran dataset yang moderat, serta karakteristik bahasa Indonesia di media sosial yang sangat beragam.

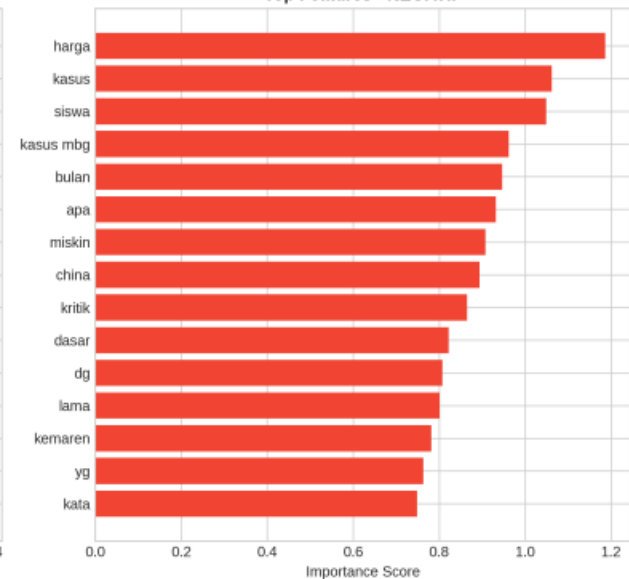
Top Features - POSITIF



Top Features - NETRAL



Top Features - NEGATIF



Keterbatasan Penelitian

- Ukuran dataset relatif terbatas (1.050 tweet) dibandingkan penelitian sejenis yang menggunakan 3.503 tweet.
- Metode pelabelan manual oleh satu anotator tunggal (tanpa multiple annotator agreement) rentan terhadap bias subjektivitas.
- Parameter SVM menggunakan default tanpa hyperparameter tuning ekstensif (mis. Grid Search).
- Skema pembagian data tunggal (single hold-out split) 80:20 tanpa validasi silang k-fold, sehingga hasil evaluasi lebih rentan terhadap variasi karakteristik satu partisi data tertentu.
- Selisih akurasi NB–SVM sebesar 1,91 poin persentase tergolong kecil dan belum diuji signifikansi statistiknya (mis. paired t-test atau McNemar's test).

Kesimpulan

Penelitian ini menunjukkan bahwa Naive Bayes memperoleh performa sedikit lebih baik dibandingkan Support Vector Machine pada seluruh metrik evaluasi serta jauh lebih efisien dari sisi waktu komputasi. Selain itu, distribusi sentimen publik terhadap Program MBG relatif seimbang, sedangkan kelas netral menjadi kategori yang paling sulit diklasifikasikan. Penelitian selanjutnya disarankan menggunakan dataset yang lebih besar, validasi silang, optimasi parameter, dan pendekatan deep learning seperti IndoBERT.

Kontribusi Penelitian

Akademis

- Memperkaya literatur analisis sentimen berbahasa Indonesia
- Memberikan perbandingan komprehensif NB vs SVM untuk isu kebijakan publik
- Menggunakan data terkini (2025) dengan distribusi sentimen seimbang

Praktis

- Rekomendasi metode terbaik untuk monitoring sentimen program pemerintah
- Dapat diaplikasikan untuk program serupa di masa depan
- Membantu evaluasi kebijakan berbasis data real-time

Terima Kasih

PERTANYAAN & DISKUSI

Mohamad Yachob Oktavian

NIM: 221080200066

Program Studi Informatika